

US012375427B2

(12) **United States Patent**
Lochhead

(10) **Patent No.: US 12,375,427 B2**
(45) **Date of Patent: Jul. 29, 2025**

(54) **NETWORK CONFIGURATION OF
TOP-OF-RACK SWITCHES ACROSS
MULTIPLE RACKS IN A DATA CENTER**

(56) **References Cited**
U.S. PATENT DOCUMENTS

(71) Applicant: **Cyxtera Data Centers, Inc.**, Coral
Gables, FL (US)

(72) Inventor: **Jason Anthony Lochhead**, Frisco, TX
(US)

(*) Notice: Subject to any disclaimer, the term of this
patent is extended or adjusted under 35
U.S.C. 154(b) by 0 days.

(21) Appl. No.: **17/847,063**

(22) Filed: **Jun. 22, 2022**

(65) **Prior Publication Data**
US 2022/0321501 A1 Oct. 6, 2022

Related U.S. Application Data

(63) Continuation of application No. 16/442,997, filed on
Jun. 17, 2019, now Pat. No. 11,374,879.

(51) **Int. Cl.**
H04L 49/25 (2022.01)
H04L 12/46 (2006.01)
H04L 41/0803 (2022.01)

(52) **U.S. Cl.**
CPC **H04L 49/25** (2013.01); **H04L 12/4641**
(2013.01); **H04L 41/0803** (2013.01)

(58) **Field of Classification Search**
CPC . H04L 49/25; H04L 12/4641; H04L 41/0803;
H04L 49/15; H04L 41/0806
See application file for complete search history.

8,458,329	B2	6/2013	Kolin et al.
8,484,355	B1	7/2013	Lochhead et al.
9,288,101	B1	3/2016	Dalal et al.
9,294,349	B2	3/2016	Jain et al.
10,142,173	B2	11/2018	Elisha
10,263,832	B1	4/2019	Ghosh
10,659,306	B2	5/2020	Okamoto et al.
10,862,821	B1 *	12/2020	Jonsson H04L 41/18
11,075,981	B2 *	7/2021	Park H04L 49/70
11,374,879	B2	6/2022	Lochhead
2004/0037295	A1	2/2004	Tanaka et al.
2006/0164245	A1	7/2006	Kibalo
2010/0142544	A1	6/2010	Chapel et al.

(Continued)

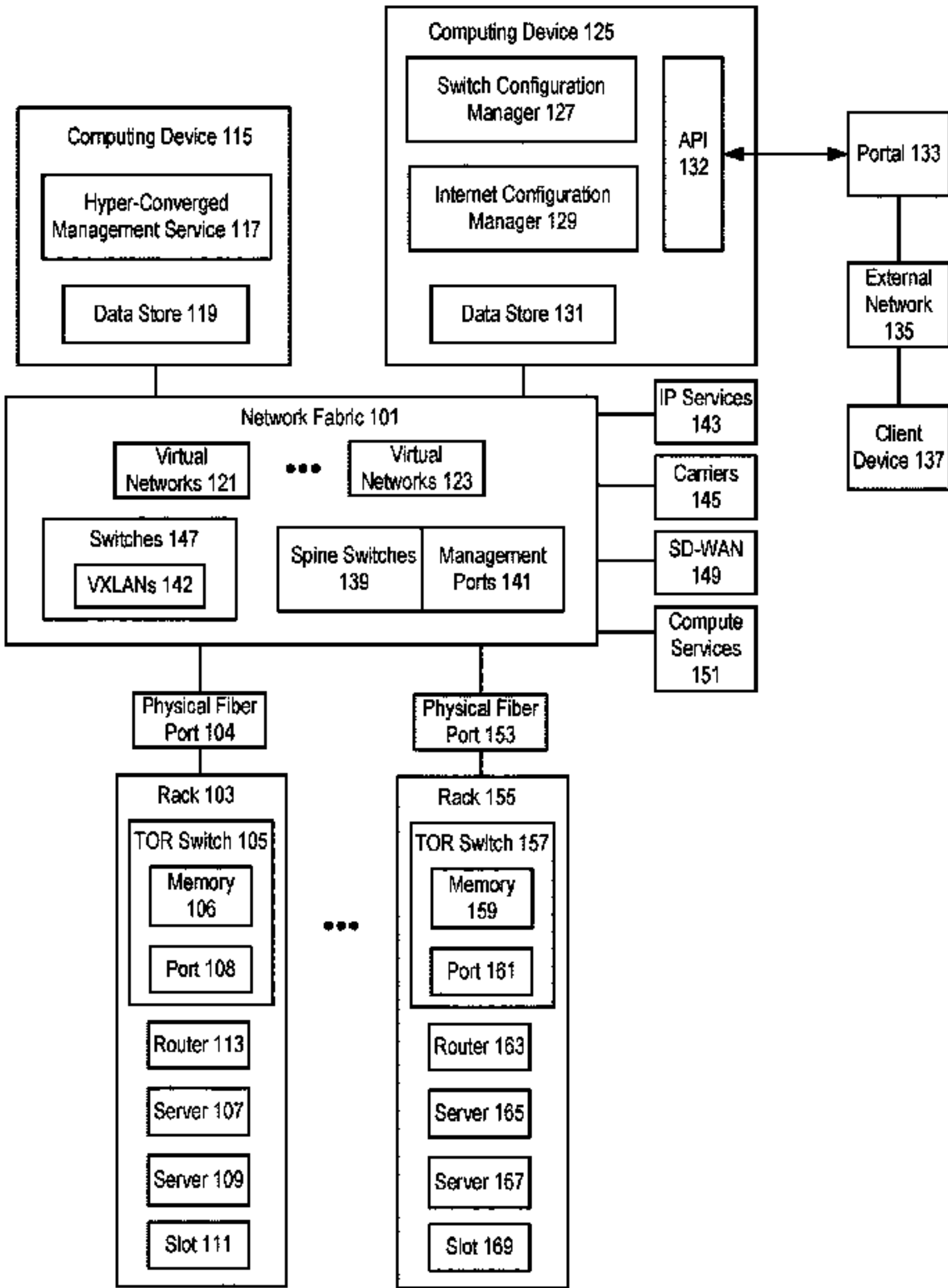
OTHER PUBLICATIONS

Network Cofiguration of Top-Of-Rack Switches Across Multiple
Racks in a Data Center, U.S. Appl. No. 16/442,997, filed Jun. 17,
2019, Jason Lochhead, Patented Case, Sep. 22, 2021.
(Continued)

Primary Examiner — Atique Ahmed
(74) *Attorney, Agent, or Firm* — Greenberg Traurig

(57) **ABSTRACT**
Top-of-rack (TOR) switches are connected to a network
fabric of a data center. Each TOR switch corresponds to a
respective rack of the data center, and is configured to
provide access to the network fabric for computing devices
mounted in the respective rack. A request is received, from
a client device via a portal, to configure a first rack of the
data center. Configuration data is received, from the client
device, for a first virtual network to be accessed by a first
computing device mounted in the first rack. In response to
receiving the configuration data, a first TOR switch of the
first rack is configured, which includes associating the first
virtual network with the first TOR switch.

16 Claims, 6 Drawing Sheets



(56)

References Cited

U.S. PATENT DOCUMENTS

2012/0110154 A1

5/2012

Adlung et al.

2013/0054426 A1

2/2013

Rowland et al.

2014/0068703 A1

3/2014

Balus et al.

2014/0351812 A1 *

11/2014

Imai H04L 12/4633
718/1

2015/0009831 A1

1/2015

Graf

2015/0103692 A1

4/2015

Jain et al.

2015/0304117 A1

10/2015

Dong et al.

2015/0379150 A1

12/2015

Duda

2016/0013974 A1

1/2016

Reddy et al.

2016/0019086 A1

1/2016

Lee et al.

2016/0087859 A1

3/2016

Kuan et al.

2016/0091685 A1

3/2016

Raza et al.

2016/0149795 A1

5/2016

Kang et al.

2017/0054603 A1

2/2017

Kulkarni et al.

2017/0093723 A1 *

3/2017

Sane H04L 47/125

2017/0149931 A1

5/2017

Lochhead et al.

2018/0145955 A1

5/2018

Nirwal

2018/0157532 A1

6/2018

Kumar et al.

2018/0176124 A1 *

6/2018

Kancherla H04L 61/2521

2018/0241629 A1

8/2018

Sato et al.

2018/0359134 A1 *

12/2018

Pech H04L 63/101

2018/0367607 A1

12/2018

Gakhar et al.

2019/0007382 A1 *

1/2019

Nirwal G06F 9/54

2019/0028342 A1 *

1/2019

Kommula H04L 45/02

2020/0028786 A1

1/2020

Chachmon et al.

2020/0136950 A1 *

4/2020

D M H04L 67/1025

2020/0344120 A1 *

10/2020

Pianigiani H04L 61/5014

2020/0382419 A1 *

12/2020

Yuan G06F 13/126

2020/0396176 A1

12/2020

Lochhead

2021/0058388 A1 *

2/2021

Knotwell G06F 9/45558

2023/0011736 A1 *

1/2023

Chitalia H04L 43/0811

OTHER PUBLICATIONS

International Search Report & Written Opinion, PCT/US2020/036131, mailed on Sep. 17, 2020.

Perry, Christian “Cyxtera: Slow provisioning derails IT transformation efforts,” 451 Research, Voice of the Enterprise: Servers and Converged Infrastructure, Budgets and Outlook, 2017, Published Apr. 6, 2018 (8 pages).

Extended European Search Report, EP20826553.8, mailed on May 25, 2023.

* cited by examiner

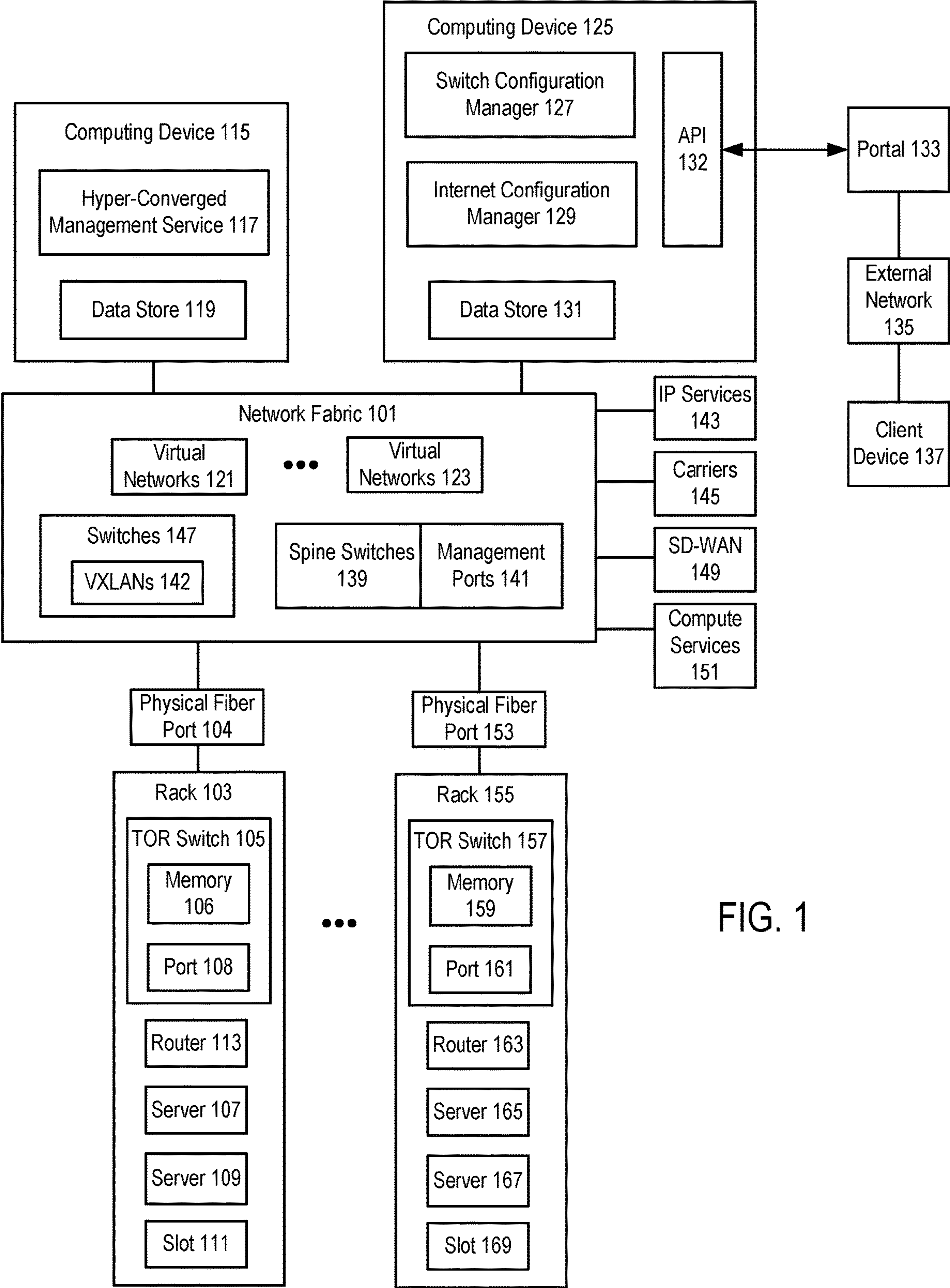
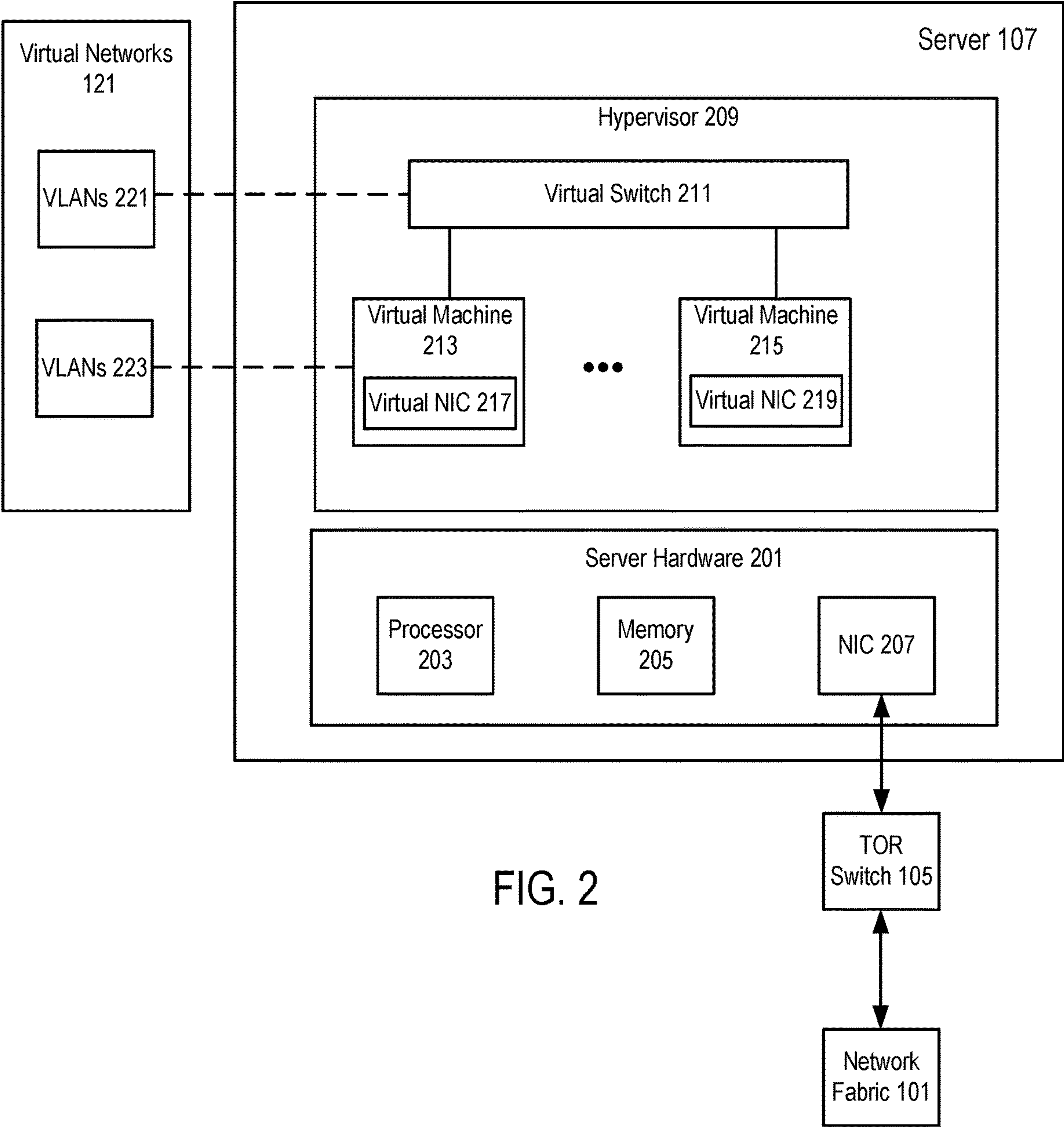


FIG. 1



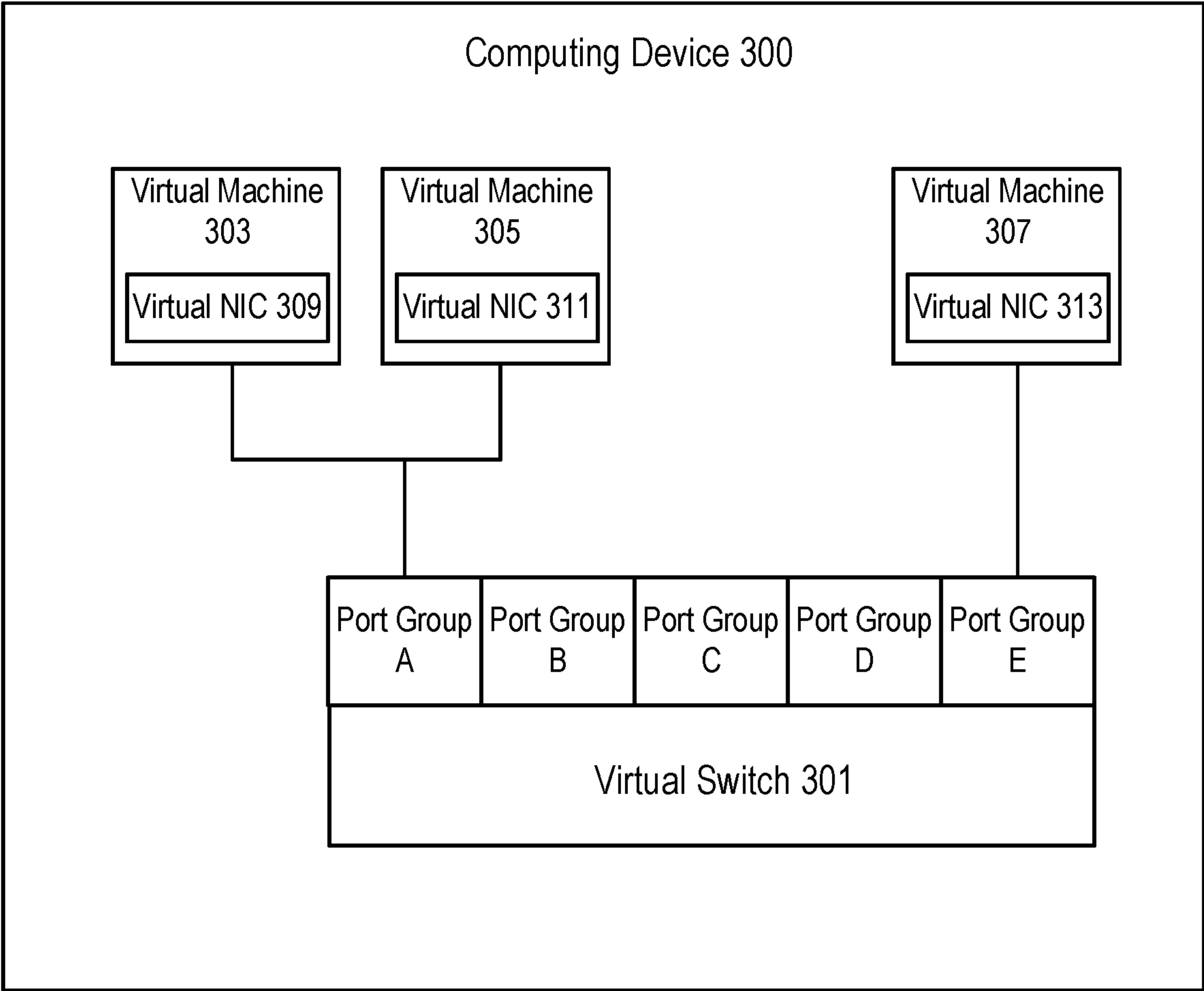


FIG. 3

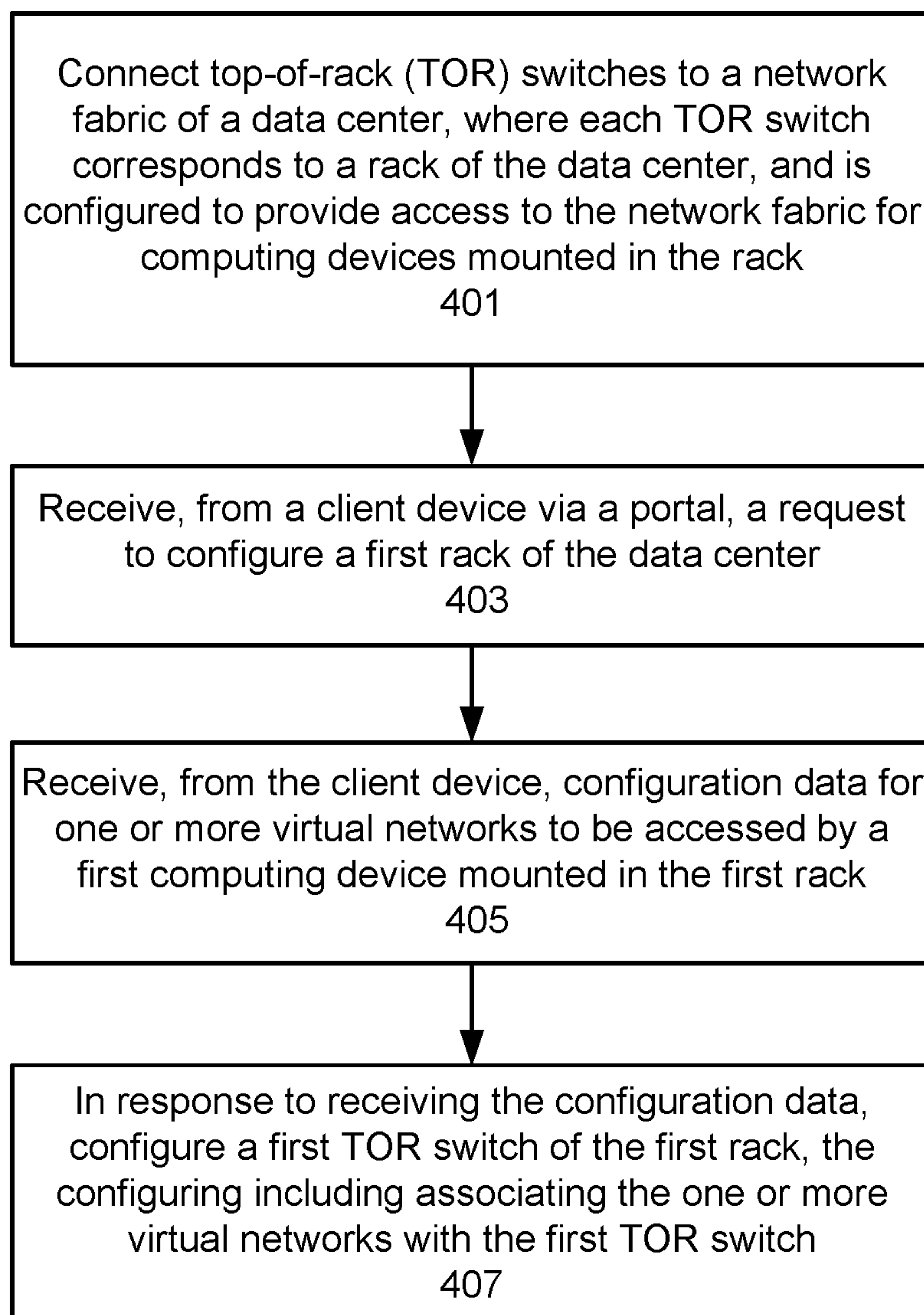


FIG. 4

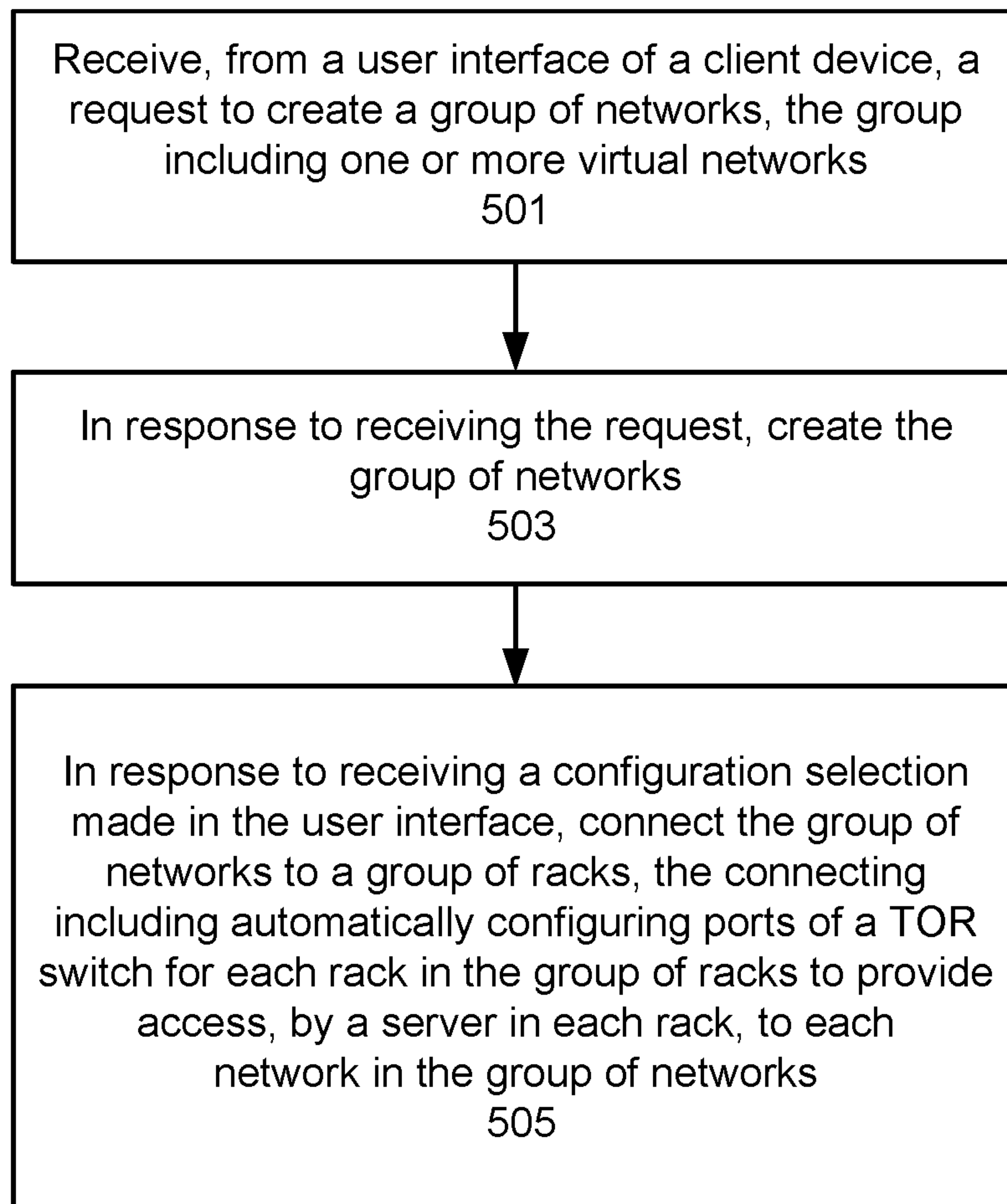


FIG. 5

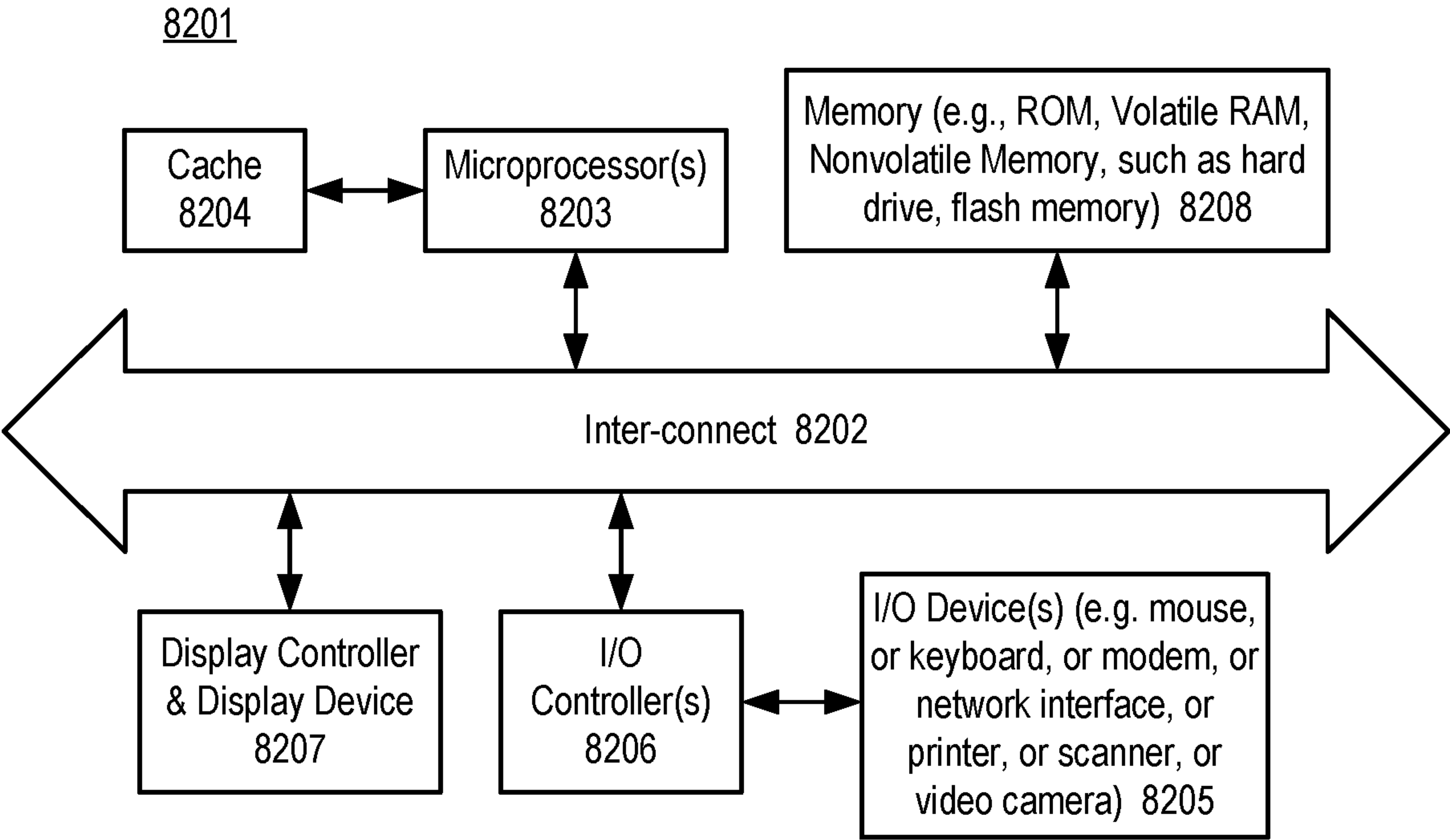


FIG. 6

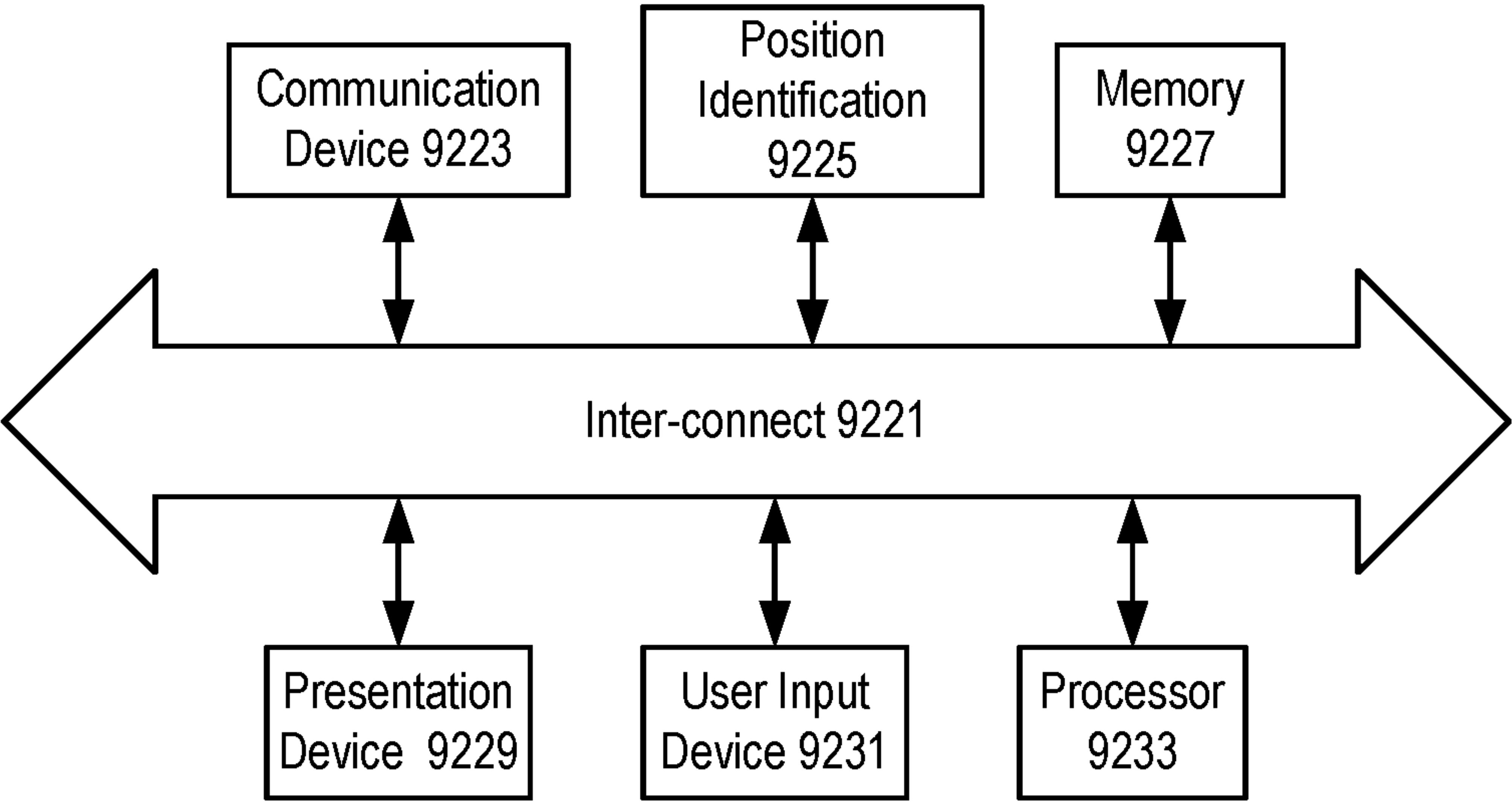


FIG. 7

1

NETWORK CONFIGURATION OF TOP-OF-RACK SWITCHES ACROSS MULTIPLE RACKS IN A DATA CENTER

RELATED APPLICATION

The present application is a continuation application of U.S. patent application Ser. No. 16/442,997, filed Jun. 17, 2019, issued as U.S. Pat. No. 11,374,879 on Jun. 28, 2022, the entire disclosure of which application is hereby incorporated herein by reference.

FIELD OF THE TECHNOLOGY

At least some embodiments disclosed herein relate to computer network configuration in general and more particularly, but not limited to network configuration for network switches in a data center.

BACKGROUND

A data center is a physical facility that houses computing systems and related networking equipment. For example, a service provider can house its computer servers at one physical location in order to manage the servers more efficiently. The servers in a data center are typically connected to users of the computer servers via the Internet, or a wide area network (WAN). The computer servers in the data center typically host applications and provide services.

The computer servers and other related components such as network switches, routers, etc., in a data center are housed in metallic cages referred to as racks. For example, a rack includes a chassis to house the computer servers. In some cases, a computer server in the form of a blade is mounted to the chassis. The rack has a wire harness for network cables that connect each blade to a computer network. Other cables provide power to each blade.

In some cases, each server mounted in the rack may be configured to host one or more virtual machines. The servers in the rack are connected to top-of-rack (TOR) switch devices. The TOR switches are connected to other TOR switches via a spine switch or spine underlay fabric. This provides a physical network that can be used by multiple tenant networks to exchange data communications between host devices in different rack units. For example, packets of data may be sent from a virtual machine in one rack unit to a virtual machine in another rack unit. The packets can be routed between corresponding TOR switch devices and an intermediary spine switch. The TOR switches are configured to store address information associated with the host devices in the data center environment.

TOR switches typically manage communications (e.g., routing and forwarding) that originate from and/or destined for physical servers (and virtual machines and virtual switches hosted by the physical servers) in a rack. Each TOR switch can be configured to communicate with a network controller unit that manages communications between TOR switches in different racks. In some cases, tenant networks residing in an underlay fabric can be created, modified, provisioned, and/or deleted.

In one example, virtual switches and virtual machines are created and run on each physical server on top of a hypervisor. Each virtual switch can be configured to manage communications of virtual machines in a particular virtual network. Each virtual machine is a member of a tenant network (e.g., a layer 3 subnet that contains one or more VLANs).

2

In one example, a TOR switch includes network ports for receiving and sending data packets to and from physical servers mounted in the racks. The ports are coupled to a switch application specific integrated circuit (ASIC) that enables packets received on one port to be forwarded to a device in the system via a different port.

In some cases, the TOR switches are used in a hyper-converged infrastructure (HCI) computing environment. HCI is a virtual computing platform used to converge computer, networking, and storage into a single software-defined architecture. The HCI computing environment can include thousands of devices such as servers and network switches. HCI services can be used to configure the network switches. In one example, an internet protocol (IP) address is configured for each network switch. In one example, an HCI management service maintains a listing of network configurations applied to network switches in various racks. In one example, the management service accesses a listing of network configurations applied to a first network switch, and dynamically applies the network configurations to a second network switch.

In one example, a first network switch resides in a slot on a rack of a data center. The HCI management service uses a data store or other memory to maintain network configurations that have been applied to the first network switch. For example, the network configurations may include switch bring-up configurations, management cluster configurations, and workload configurations.

BRIEF DESCRIPTION OF THE DRAWINGS

The embodiments are illustrated by way of example and not limitation in the figures of the accompanying drawings in which like references indicate similar elements.

FIG. 1 shows an example data center that includes a network fabric connecting top-of-rack switches for racks in which various computing equipment is mounted, according to one embodiment.

FIG. 2 shows an example server including server hardware executing a hypervisor that supports virtual machines, according to one embodiment.

FIG. 3 shows an example computing device running virtual machines that connect to ports of a virtual switch, according to one embodiment.

FIG. 4 shows a method for configuring a top-of-rack switch that is connected to a network fabric of a data center, according to one embodiment.

FIG. 5 shows a method for connecting a group of networks to a group of racks in response to a configuration selection received from a user by input into a user interface, according to one embodiment.

FIG. 6 shows a block diagram of a computing device, which can be used in various embodiments.

FIG. 7 shows a block diagram of a computing device, according to one embodiment.

DETAILED DESCRIPTION

At least some embodiments herein relate to configuring network switches in a data center. In one example, the network switches are top-of-rack (TOR) switches. Alternatively and/or additionally, other types of network switches can be configured. In one example, the TOR switches are connected to a network fabric of the data center. The network fabric connects TOR switches used in various racks that are housed in the data center. Each rack mounts various computing hardware such as physical servers, routers, etc.

A significant problem with existing data centers is the long time required to deploy information technology (IT) infrastructure in a data center (e.g., provisioning of new hardware servers and applications that run on the servers). For example, it can take three to six months to deploy a single application, including provisioning of circuits, building out infrastructure in a colocation cage, installation and configuration of the hypervisor, and loading and testing of the application. The time to provision network connectivity and services often constrains colocation deployments of new workloads or applications.

Another problem is difficulty in accurately forecasting bandwidth and overall IT capacity requirements more than a few months in advance. This results in many organizations initially over-provisioning to assure that adequate bandwidth and compute resources are available as demand grows.

Various embodiments described below provide a technological solution to one or more of the above technical problems. In one embodiment, a method includes: mounting a switch in a rack (e.g., a TOR switch of a rack in a data center), wherein the rack is configured for mounting a server connected to the switch; connecting the switch to a network fabric; receiving, by a switch configuration manager from a client device (e.g., a client device of a service provider customer that is deploying new IT infrastructure in the data center), instructions to create a virtual network; in response to receiving the instructions, creating the virtual network; and configuring, by the switch configuration manager and based on the instructions, the switch to associate the virtual network with the switch.

In one example, the switch configuration manager is software executed by a computing device connected to the network fabric of a data center that houses racks of computer hardware, including the rack above. The switch configuration manager configures the TOR switches for all racks physically located in the data center. The switch configuration manager is accessed by service provider customers of the data center using an application programming interface (API) of the switch configuration manager. For example, a client device for each customer can use the API to configure the TOR switches for its racks when the customer is deploying new IT infrastructure in the data center.

In one case, the virtual network is a first virtual network, and the method further includes: receiving, from a user interface of the client device, a request to create a group of networks, the group including the first virtual network; in response to receiving the request, creating the group of networks; and in response to receiving a user selection made in the user interface, connecting the group of networks to a group of racks, the connecting including automatically configuring ports of a TOR switch for each rack in the group of racks to provide access, by a respective server in each rack, to each network in the group of networks.

In one embodiment, a service provider or other customer is provided a user interface (UI) and an API. In one example, the service provider is a cloud service provider, a software as a service (SaaS) provider, or a managed hosting provider. The UI presents customer ports, compute nodes, and other elements connected to the network fabric of the data center. The customer can create virtual networks or groups of virtual networks using the UI. The customer can bundle several virtual networks into a defined group (and optionally assign a text label to the group). The customer can then use the UI to connect the defined group between racks and other computing devices. Data center automation software (e.g., executing on a virtual server of the data center) examines

data for the group and configures connections for the virtual networks in the group as needed.

In one embodiment, the data center automation software manages network connections to a customer's racks. The customer can use a portal (e.g., provided by a user application executing on a client device such as a mobile device) to connect a group of networks to a group of racks. In one example, each rack has a unique ID. The customer can see rack data, including location by metro region, on a display of its client device using the UI. The customer can also see IP connectivity instances (e.g., by metro or other geographic region) and ports in a metro or other geographic region that can be used to receive services over the network fabric. For example, multiple racks can all access the same IP connectivity instance. The portal displays endpoints and connections on the customer's client device, and the portal manages the relationship between the endpoints and connections.

In one embodiment, the portal provides control by the customer of a grouping mechanism for the customer's racks. The customer can manage network connections to its racks. In one example, the customer requests that a group of networks be connected to a group of racks. Then, the data center automation software configures these connections. In one embodiment, the data center automation software includes the switch configuration manager described above.

In one example, a customer creates a group, and assigns networks to the group. On the backend, the association of the networks to the group is tracked by the data center automation software. When a new request is made to make a new connection to the group, all networks that are part of the group are examined, and individual configurations are implemented as required to make the new connection.

In one embodiment, devices and ports to be connected to the above networks are identified. For each network, endpoints are determined, and work required to implement the connections is identified as one or more workflows. In one example, a workflow engine (e.g., software executing on a virtual machine of an administrator computing device of the data center) executes tasks in the workflows.

In one embodiment, colocation racks can be delivered to the customer faster than when using prior approaches. The racks are standalone racks that include power, a locking mechanism for each rack, and network switches that are tied to the network fabric of the data center. In some cases, IP transit is provided for servers in the racks for internet connectivity. In one example, a customer signs a service agreement and the customer is added to an authentication service used in the data center. The authentication service manages access by and identifies the customer for the data center. The customer logs into a command center of the data center (e.g., the command center can be implemented by software that includes the switch configuration manager above). The customer then selects a data center location, and specifies an order for a quantity of racks (e.g., from one rack to a predetermined limit).

Next, the command center performs various actions. The command center maintains a database of available rack inventory at various geographic data center locations worldwide. In response to the customer selection, the command center allocates racks from available inventory in the location selected by the customer. The authentication service is updated with rack assignment information corresponding to these allocated racks.

A security system at each physical data center facility where the selected racks are located is updated so that the customer is allowed to physically access the racks. A lock system used on the racks is configured to allow the customer

5

to access the selected racks. In some cases, IP connectivity (e.g., to provide internet access) is provision and connected to the TOR switches for the selected racks. Then, the portal is updated with the locations of the selected racks, TOR switch information for the racks, and IP connectivity information (e.g., VLAN, subnet, and default gateway configuration information) for the racks. Billing of the customer for the colocation service is initiated (e.g., by electronic communication). Finally, the customer is notified by electronic communication or otherwise when the foregoing provisioning is complete.

After the command center performs the actions above, the customer can perform various actions. First, the customer accesses the command center to complete user setup, including uploading or taking a photo via the portal. In one example, the customer accesses the command center using the client device above. In one example, the client device is a mobile device having a camera and is used to take a photo of personnel associated with the customer. The photo is uploaded to the command center via the API above.

When the customer physically arrives at a data center location, the customer checks in with security to receive a badge. The badge includes the photo previously provided by the customer above. The customer enters the facility and unlocks the selected racks using the badge. In one example, the badge contains security credentials necessary to unlock the locking mechanism on the selected racks. The customer installs computing equipment in the selected racks, and then cables the equipment to the TOR switches above. The customer then accesses the command center and configures ports of the TOR switches. In one example, the switch ports are configured with a virtual local area network (VLAN) configuration desired for use by the customer.

In one embodiment, the operator of the data center buys hardware equipment and installs it in racks. The equipment is made available to customers on demand. This permits customers to avoid having to build equipment for peak demand. In one example, a customer can purchase computing resources that are supported by this equipment. In one example, the purchased computing resources are based on a hyper-converged infrastructure (HCI). For example, the customer can use the portal above to select computing resources. In one example, the computing resources are connected to one or more virtual networks configured by the customer using the portal. The command center above configures the TOR switches to connect these virtual networks to the hardware equipment of the data center.

In one embodiment, an on-demand IT infrastructure is provided to customers. In one example, the infrastructure is provided using an on-demand consumption model. In one example, the infrastructure is a physically-isolated on-demand hyper-converged infrastructure. The network fabric is a software-defined network fabric that provides connectivity via a secure layer 2 network throughout the data center. The customer can request access to network providers with direct connections to private or public cloud resources.

In one embodiment, a customer installs its own equipment in a first rack. The customer configures the TOR switches of the first rack using a portal as described above. The command center above configures ports of the TOR switches to implement the configuration requested by the customer. In addition, the customer can configure and deploy equipment in a second rack that has been pre-installed and is owned by the operator of the data center. The second rack includes equipment that provides a so-called "compute node" for deployment by the customer. In one example, the compute node is a dedicated self-contained HCI unit that combines

6

computer resources (e.g., CPU cores), memory resources (e.g., RAM), and storage resources (e.g., hard disk drive and solid-state disk) into a pre-configured integrated appliance. A group of compute nodes forms a cluster.

In one example, the compute nodes provide dedicated hardware for a customer upon which the customer can deploy its desired hypervisor. The customer can then configure and manage the resources and virtual machines needed to run desired workloads. In one example, the customer uses the portal above to create one or more virtual networks that connect one or more servers of the first rack to one or more servers of the second rack. The first rack and second rack can be in different data centers.

In one embodiment, the network fabric of the data center above is a softwaredefined network fabric to link customers and resources throughout the data center. The network fabric uses an architecture to assure that each customers traffic is logically isolated and protected through the use of a virtual extensible local area network (VLAN) protocol. Using the API above, the client device of the customer can define, provision, and configure private virtual layer 2 networks. In one example, logical services are delivered to servers in a rack of the customer as virtual networks using VXLANs. In one example, all physical connections are delivered with an Ethernet layer 2 interface. In one example, multiple services are delivered to customer servers over a single physical connection. In one example, the physical connection is a physical port implemented using single-mode fiber operating at 1-10 Gbps.

In light of the above, automated configuration of network switches in a data center can provide one or more various advantages. For example, customer colocation access can be automated and provided more quickly than using prior approaches. For example, colocation access can be provided in less than 48 hours (e.g., the same day) from receipt of the initial request by the customer. For example, deployment of Internet connectivity to rack switches can be automated. For example, multiple security systems and multiple rack switches can be configured simultaneously. For example, self-service configuration of TOR switches across multiple racks can be provided.

Other advantages can include, for example, one or more of the following:

- Provisioning hyper-converged infrastructure in a shorter period of time.

- Extending colocation environments and connectivity within or across data centers in different geographic locations.

- Retaining full control by the customer over its network and compute environment, with dedicated hardware.

- Reducing complexity by delivering multiple services over a single physical network connection.

FIG. 1 shows an example data center that includes a network fabric **101** connecting top-of-rack (TOR) switches **105**, **157** for racks **103**, **155** in which various computing equipment is mounted, according to one embodiment. The computing equipment mounted in rack **103** includes the TOR switch **105**, and also servers **107**, **109**, and router **113**. Rack **103** has a slot **111** in which additional equipment can be mounted (e.g., slot **111** and/or other slots can be used by a customer of the data center to install customer-owned equipment in rack **103**).

TOR switch **105** includes memory **106** and various ports (e.g., port **108**) for receiving and sending communications (e.g., data packets). Memory **106** stores a network configuration (e.g., port connection assignments) as implemented by switch configuration manager **127** over network fabric

101 in response to a customer request received over a portal 133. Various ports of TOR switch 105 connect to router 113 and/or servers 107, 109. Other ports of TOR switch 105 connect to one or more virtual networks 121, 123 of network fabric 101. In one embodiment, all communications between rack 103 and network fabric 101 pass through a physical fiber port 104 (e.g., implemented using single-mode fiber).

Rack 155 mounts computer equipment including the TOR switch 157, servers 165, 167, and router 163. Rack 155 includes a slot 169 for adding additional equipment. TOR switch 157 includes memory 159 and various ports, including port 161. Similarly as for rack 103, all communications to and from the network fabric 101 pass through a physical fiber port 153. Also, similarly as for rack 103, memory 159 is used to store data regarding a configuration of TOR switch 157 as automatically implemented by switch configuration manager 127. In one example, this configuration is implemented in response to a selection made by a customer in a user interface of client device 137. The data center of FIG. 1 can include numerous other racks connected to network fabric 101 using physical fiber ports and/or other types of connections.

The virtual networks 121, 123 of network fabric 101 can overlay various types of physical network switches. In one embodiment, network fabric 101 comprises network switches 147 that are used to implement virtual extensible local area networks (VXLANs) 142 for transmission of data from a server of rack 103 to a server mounted in a different rack, such as rack 155. In one example, a virtual network connected to TOR switch 105 is converted into a VXLAN 142 for transmission of data from server 107 to server 165. The VXLAN 142 is used to transmit the data to another virtual network connected to TOR switch 157. VXLANs 142 can be configured by switch configuration manager 127 to implement the foregoing connection between servers. In one embodiment, this configuration is implemented in response to a request from client device 137 to add server 165 to a virtual network that includes server 107.

In one embodiment, network fabric 101 includes spine switches 139 as part of a physical switching fabric. Spine switches 139 include management ports 141, which can be used by switch configuration manager 127 to configure spine switches 139.

In one example, network fabric 101 is a leaf-spine data center switching fabric. In one example, network fabric 101 is a software-defined network (SDN) controller-based data center switching fabric. In one example, the switching fabric supports all workloads (e.g., physical, virtual machine, and container) and choice of orchestration software. The switching fabric provides layer 2 (L2) switching, and layer 3 (L3) routing. In one example, the switching fabric is scalable, resilient, has no single point of failure, and/or supports headless mode operations.

In one embodiment, a computing device 115 (e.g., a server or virtual machine) is connected to network fabric 101. Computing device 115 executes a hyperconverged management service 117, which can be used to allocate compute, memory, and/or storage resources provided by various racks, including rack 103 and/or rack 155. Data store 119 is used to store data regarding this allocation of resources.

In one embodiment, a customer installs its own equipment into rack 103. Using client device 137, the customer sends a request for additional resources to add to its computing environment in the data center. In response to this request, hyperconverged management service 117 allocates resources of servers in rack 155 for use by the customer. In

one example, virtual machines are created on rack 155 for handling workloads of the customer.

In one embodiment, a computing device 125 is connected to network fabric 101. Switch configuration manager 127 executes on computing device 125 and performs various administrative functions for the data center (e.g., functions as described above). Some of the functions performed by switch integration manager 127 are responsive to communications received from client device 137 over an external network 135 through portal 133. Client device 137 uses API 132 of switch configuration manager 127 for these communications. Client device 137 also receives communications from switch configuration manager 127 using API 132. In one example, one or more of the communications cause a display of information in a user interface of client device 137. In one example, the user interface uses the information to display a configuration of a computing environment of a customer of the data center.

In one embodiment, in response to a communication from client device 137, switch configuration manager 127 creates and/or configures various virtual networks of network fabric 101 (e.g., virtual networks 121, 123, and/or VXLANs 142). In one example, certain virtual networks are assigned to a group as designated by a customer using client device 137. Data regarding creation and/or configuration of virtual networks (e.g., assignment of virtual networks to a group(s)) is stored in data store 131.

In one embodiment, a customer of the data center can use client device 137 to request internet connectivity for one or more racks in its computing environment. For example, the customer can request that internet connectivity be provided for use by servers 107, 109. Communications with client device 137 regarding internet connectivity also can be performed using API 132. In response to this request, internet configuration manager 129 can configure IP services 143 to provide this internet connectivity. Internet configuration manager 129 communicates configuration data needed by switch configuration manager 127 for configuring TOR switch 105 so that servers 107, 109 are connected to IP services 143, which provides the internet connectivity. Configuration data regarding this internet connectivity can also be stored in data store 131.

In one embodiment, the customer can request that one or more telecommunications carriers 145 be connected to racks in its computing environment (e.g., rack 103 or rack 155).

In one embodiment, the customer can request that servers in rack 103 or rack 155 be connected to a software-defined wide area network (SD-WAN) 149. In one example, SD-WAN 149 is used by a customer to extend its computer networks over large distances, to connect remote branch offices to data centers and each other, and/or to deliver applications and services required to perform various business functions.

In one embodiment, the customer can request compute services 151. In one example, compute services 151 include one or more virtual machines created for use in the customer's computing environment. In one example, the virtual machines are created and run on servers in racks of the data center. For example, hyper-converged management service 117 can create and manage these virtual machines.

In another example, compute services 151 include storage resources. The storage resources can be non-volatile memory devices mounted in racks of the data center (e.g., mounted in rack 155).

In one embodiment, a virtualization control system (e.g., implemented by hyper-converged management service 117 or otherwise by computing device 115) abstracts server,

storage, and network hardware resources of the data center to provide a more granular virtual server, virtual storage, and virtual network resource allocation that can be accessed by a customer. A customer console provisioning interface is coupled to the virtualization control system to permit the customer to configure its new environment. In one example, the virtualization control system responds to requests received from client device 137.

In one embodiment, portal 133 is a web portal. Client device 137 provides a user interface that enables a customer/user to associate a specified network connection with a new computing environment. The new computing environment can be associated with a number of virtual machines that is specified in the user interface.

In one embodiment, a customer can use the user interface to create, provision, and manage its virtual resources across numerous virtual environments (which may physically span multiple physical data centers). For example, some virtual servers are physically located on hardware in a first physical data center, and other virtual servers are physically located in a second physical data center. In one example, the difference in physical location is irrelevant to the customer because the customer is presented an abstracted view of data center assets that span multiple virtualization control systems and multiple geographic locations.

In one embodiment, the above user interface enables a customer/user to add a network to a newly-created environment. The network is given a name and a VLAN identifier. The customer can create and place a new virtual server within the new environment. The customer can configure processing, memory, and storage resources to be associated with the new virtual server being created. The new server can then be deployed to the customer environment.

In one embodiment, the customer uses the user interface to perform configuration tasks for the new virtual server (e.g., providing a server name, selecting a number of processors to be associated with the virtual server, selecting an amount of system memory to be associated with the virtual server). The customer selects an operating system to associate with the new server.

In one embodiment, a customer can create groups of virtual servers. For example, customers can organize servers by function (e.g., a group of web servers, a group of SQL servers). The customer selects a particular virtual network (e.g., virtual network 121) to associate with the virtual server (e.g., a virtual machine running on server 107 or server 165), and then provides details of the IP address and DNS settings for the virtual server.

In one embodiment, after a customer purchases a block of IP addresses (e.g., associated with IP services 143), public IP addresses can be displayed in the user interface on client device 137. Another display screen can allow a user to examine assignments of private IPs to different virtual servers that have been configured.

In one embodiment, the user interface on client device 137 can be used to create an Internet service. The user selects a public IP address and a protocol. The user may then select a port value and a service name. A service description may be provided. A list of Internet services that have been provisioned for the IP address can be displayed in the interface. The provisioned services can include, for example, an FTP service, an SMTP service, etc. Within each service are listed the nodes (e.g., virtual servers) that have been created and associated with a particular Internet service, as well as the protocol and port.

In one example, switch configuration manager 127 can access the above customer environments (e.g., to add a network to a customer environment).

FIG. 2 shows server 107 of FIG. 1, according to one embodiment. Server 107 includes server hardware 201 that executes a hypervisor 209. The hypervisor 209 supports virtual machines 213, 215. The server hardware 201 includes a processor 203, memory 205, and a network interface controller (NIC) 207. NIC 207 connects server 107 to a port of TOR switch 105. Another port of TOR switch 105 is connected to network fabric 101.

Virtual machines 213, 215 generally communicate with network fabric 101 using TOR switch 105. Virtual machine 213 has a virtual NIC 217, and virtual machine 215 has a virtual NIC 219. In one embodiment, virtual NICs 217, 219 connect virtual machines 213, 215 to one or more virtual networks 121 of network fabric 101. In one example, virtual machine 213 is associated with VLANs 223 of network fabric 101.

For example, VLANs 223 may have been created by a customer of the data center that itself has installed server 107 in rack 103. In one example, the customer installs server 107 after switch configuration manager 127 has configured one or more ports of TOR switch 105 in response to one or more communications from client device 137. In one example, a locking mechanism on rack 103 does not permit entry by the customer until this configuration of TOR switch 105 has been completed by switch configuration manager 127.

Hypervisor 209 also supports a virtual switch 211. Virtual machines 213, 215 are connected to ports of virtual switch 211. In one example, virtual switch 211 also has one or more ports associated with VLANs 221 of network fabric 101.

FIG. 3 shows an example computing device 300 running virtual machines 303, 305, 307 that connect to various ports of a virtual switch 301, according to one embodiment. Computing device 300 is an example of server 107 of FIG. 2.

The ports of virtual switch 301 are provided in various groups (e.g., Port Group A, B, C, D, E). In one example, virtual machines 303, 305 are connected to Port Group A via virtual NICs 309, 311. Virtual machine 307 is connected to Port Group E via virtual NIC 313.

In one example, each port group corresponds to a virtual network. Virtual switch 300 is an example of virtual switch 211 of FIG. 2. In one example, each port group corresponds to one of VLANs 223 of FIG. 2.

In one case, computing device 300 is an example of computing device 125 of FIG. 1. Switch configuration manager 127 and/or internet configuration manager 129 can be implemented using virtual machines 303, 305, and/or 307. In one example, virtual machine 307 is used to implement portal 133 for communications with client device 137 using API 132.

In one embodiment, computing device 300 is used to implement compute services 151 of FIG. 1. For example, a customer can use client device 137 to request that one or more of virtual machines 303, 305, 307 be allocated to the customer's computing environment. In response to this request, one or more virtual networks of the customer's computing environment are connected to one or more of virtual machines 303, 305, 307. In one example, Port Group A corresponds to a group of virtual machines requested by the customer using client device 137. In one example, Port Group A corresponds to a group that is created in response to a customer request.

FIG. 4 shows a method for configuring a top-of-rack switch (e.g., TOR switch 105 of FIG. 1) that is connected to

11

a network fabric (e.g., network fabric **101** of the data center of FIG. **1**), according to one embodiment. For example, the method of FIG. **4** can be implemented in the system of FIGS. **1**, **2**, and **3**.

The method of FIG. **4** can be performed by processing logic that can include hardware (e.g., processing device, circuitry, dedicated logic, programmable logic, microcode, hardware of a device, integrated circuit, etc.), software (e.g., instructions run or executed on a processing device), or a combination thereof.

In some embodiments, the method of FIG. **4** is performed at least in part by one or more processors of computing device **125** of FIG. **1**. In one embodiment, computing device **125** is implemented using the processors and memory of FIG. **6** or **7** (see below).

Although shown in a particular sequence or order, unless otherwise specified, the order of the processes can be modified. Thus, the illustrated embodiments should be understood only as examples, and the illustrated processes can be performed in a different order, and some processes can be performed in parallel. Additionally, one or more processes can be omitted in various embodiments. Thus, not all processes are required in every embodiment. Other process flows are possible.

At block **401**, TOR switches are connected to a network fabric of a data center. Each TOR switch corresponds to a rack of the data center, and is configured to provide access to the network fabric for one or more computing devices mounted in the rack. In one example, the TOR switches are TOR switches **105** and **157** of FIG. **1**, and the network fabric is network fabric **101**. In one example, the computing devices mounted in the rack include router **113** and server **107**, which are mounted by a customer after the rack has been assigned to the customer. This rack assignment occurs after the customer has requested the rack using client device **137**.

At block **403**, a request is received from a client device via a portal. The request is to configure a first rack of the data center. In one example, the client device is client device **137**, and the portal is portal **133**.

At block **405**, configuration data is received from the client device. The configuration data is for one or more virtual networks to be accessed by a first computing device mounted in the first rack. In one example, the configuration data includes a specification of the devices and ports that a customer desires to connect to each of the virtual networks. In one example, the configuration data includes IP addresses associated with internet connectivity (e.g., provided by IP services **143**). In one example, the configuration data includes a subnet mask and an identification of a gateway (e.g., for use in configuring a router). In one example, the virtual networks include virtual networks **121** and **123** of FIG. **1**.

At block **407**, in response to receiving the configuration data, a first TOR switch of the first rack is configured. This configuration includes associating the one or more virtual networks with the first TOR switch. In one example, switch configuration manager **127** configures TOR switch **105** of rack **103**. This configuration includes associating virtual networks **121** with TOR switch **105**.

In one embodiment, a method comprises: mounting a switch (e.g., TOR switch **105**) in a rack (e.g., rack **103**), wherein the rack is configured for mounting a server (e.g., server **107**) connected to the switch; connecting the switch to a network fabric (e.g., network fabric **101**); receiving, by a switch configuration manager (e.g., switch configuration manager **127**) from a client device (e.g., client device **137**),

12

instructions to create a virtual network (e.g., one of virtual networks **121**); in response to receiving the instructions, creating the virtual network; and configuring, by the switch configuration manager and based on the instructions, the switch to associate the virtual network with the switch.

In one embodiment, the method further comprises converting the virtual network into a virtual extensible local area network (e.g., one of VXLANs **142**) for transmission of data from the server over the network fabric to a server mounted in a different rack.

In one embodiment, the rack is a first rack, the server is a first server, and the switch is a first switch. The method further comprises: receiving, by the switch configuration manager from the client device, instructions to associate the virtual network with a second server mounted in a second rack (e.g., rack **155**); and in response to receiving the instructions to associate the virtual network with the second server, configuring a second switch (e.g., TOR switch **157**) of the second rack to associate the VXLAN with the second switch.

In one embodiment, the virtual network is a first virtual network, and the method further comprises: receiving, from the client device, an instruction to create a second virtual network associated with the second server; and in response to receiving the instruction to create the second virtual network, configuring the network fabric to associate the second virtual network with the second server.

In one embodiment, the virtual network is a first virtual network, and the method further comprises: receiving, from the client device, an instruction to create a group including the first virtual network and a second virtual network; in response to receiving the instruction to create the group, storing data regarding the group in a data store (e.g., data store **131**) that stores configuration data for switches in the network fabric; receiving, from the client device, an instruction to connect a virtual server to the group; and in response to receiving the instruction to connect the virtual server to the group, configuring at least one switch of the network fabric to associate the virtual server with the first virtual network and the second virtual network.

In one embodiment, the rack is a first rack in a first data center at a first geographic location, and the virtual network is a first virtual network. The method further comprises: receiving, from the client device, an instruction to create a second virtual network; in response to receiving the instruction to create the second virtual network, configuring the network fabric to create the second virtual network; receiving an instruction to create a group including the first virtual network and the second virtual network; in response to receiving the instruction to create the group, updating, by the switch configuration manager, a data store (e.g., data store **131**) to track membership of the first virtual network and the second virtual network in the group; receiving, from the client device, an instruction to connect the group to a second rack in a second data center at a second geographic location; and in response to receiving the instruction to connect the group to the second rack, configuring the network fabric to associate the second virtual network with a switch of the second rack.

In one embodiment, a method comprises: connecting top-of-rack (TOR) switches to a network fabric of at least one data center (e.g., the data center of FIG. **1**), wherein each TOR switch corresponds to a respective rack of the at least one data center, and is configured to provide access to the network fabric for computing devices mounted in the respective rack; receiving, from a client device via a portal (e.g., portal **133**), a request to configure a first rack of the at

13

least one data center; receiving, from the client device, configuration data for at least one first virtual network to be accessed by a first computing device mounted in the first rack; and in response to receiving the configuration data, configuring a first TOR switch of the first rack, the configuring including associating the at least one first virtual network with the first TOR switch.

In one embodiment, the computing devices are physical servers (e.g., server **107** of FIG. **2**) configured to run virtual servers (e.g., virtual machines **213**, **215**), and the physical servers include a first physical server configured to run a first virtual server, the method further comprising configuring a virtual extensible local area network (VXLAN) of the network fabric to connect the first TOR switch to a second TOR switch of a second rack of the at least one data center, wherein the VXLAN is configured to transmit data from the first virtual server to a second virtual server running on a second physical server mounted in the second rack.

In one embodiment, each of the computing devices is a physical server, a network device, or a storage device; and the first TOR switch comprises at least one port, and configuring the first TOR switch comprises configuring the at least one port based on the configuration data.

In one embodiment, the first rack comprises a second TOR switch. A first port of the first TOR switch and a second port of the second TOR switch are configured for connection to the first computing device.

In one embodiment, the first TOR switch comprises a port, and configuring the first TOR switch comprises associating a virtual local area network (VLAN) with the port.

In one embodiment, the method further comprises: causing display, in a user interface of the client device, of an identifier for the first rack, and a geographic location of the first rack, wherein the identifier for the first rack is stored in a data store, and wherein the user interface enables a user to request that at least one virtual network be created in the network fabric; and storing, in the data store, a name and an identifier for each of the created at least one virtual network.

In one embodiment, the method further comprises causing display, in a user interface of the client device, of availability of ports for each of a plurality of geographic locations in which racks, including the first rack, are located, wherein each of the ports provides a connection to at least one of IP services (e.g., IP services **143**) or compute services (e.g., compute services **151**) over the network fabric.

In one embodiment, configuring the first TOR switch further includes providing access for the first computing device to the IP services or compute services.

In one embodiment, the client device generates the configuration data based on inputs received by a user interface of the client device. The inputs include selection of an icon in the user interface that corresponds to the first rack, and selection of the icon causes presentation in the user interface of configuration options for the first TOR switch.

In one embodiment, the first computing device has a port configured to connect to the at least one virtual network.

FIG. **5** shows a method for connecting a group of networks to a group of racks in response to a configuration selection received from a user by input into a user interface, according to one embodiment. For example, the method of FIG. **5** can be implemented in the system of FIGS. **1**, **2**, and **3**.

The method of FIG. **5** can be performed by processing logic that can include hardware (e.g., processing device, circuitry, dedicated logic, programmable logic, microcode,

14

hardware of a device, integrated circuit, etc.), software (e.g., instructions run or executed on a processing device), or a combination thereof.

In some embodiments, the method of FIG. **5** is performed at least in part by one or more processors of computing device **125** of FIG. **1**. In one embodiment, computing device **125** is implemented using the processors and memory of FIG. **6** or **7** (see below).

Although shown in a particular sequence or order, unless otherwise specified, the order of the processes can be modified. Thus, the illustrated embodiments should be understood only as examples, and the illustrated processes can be performed in a different order, and some processes can be performed in parallel. Additionally, one or more processes can be omitted in various embodiments. Thus, not all processes are required in every embodiment. Other process flows are possible.

At block **501**, a request is received from a client device. The request is based on input provided into a user interface of the client device. The request is to create a group of networks, where the group includes one or more virtual networks. In one example, the client device is client device **137**. In one example, the input is provided by a customer of the data center. The customer may provide a name which is assigned to the group.

At block **503**, in response to receiving the request, the group of networks is created. In one example, the group of networks is created by switch configuration manager **127**. The virtual networks that are assigned to the group are stored in data store **131**.

At block **505**, in response to receiving a configuration selection made in the user interface, the group of networks is connected to a group of racks. The connecting includes automatically configuring ports of a TOR switch for each rack in the group racks to provide access, by a server of each rack, to each network in the group of networks. In one example, the group of networks is connected to the group of racks by switch configuration manager **127**. For example, network fabric **101** and TOR switches **105** and **157** are configured to connect each network of the group to racks **103** and **155**. In one example, the group of networks includes virtual networks **121** and/or **123**.

In one embodiment, a method comprises: receiving, by a switch configuration manager (e.g., switch configuration manager **127**) from a client device (e.g., client device **137**), instructions to create a virtual network (e.g., one of virtual networks **121**); in response to receiving the instructions, creating the virtual network; and configuring, by the switch configuration manager and based on the instructions, a switch (e.g., TOR switch **105**) to associate the virtual network with the switch.

In one embodiment, the virtual network is a first virtual network, and the method further comprises: receiving, from a user interface of the client device, a request to create a group of networks, the group including the first virtual network; in response to receiving the request, creating the group of networks; and in response to receiving a user selection made in the user interface, connecting the group of networks to a group of racks (e.g., racks **105** and **157**), the connecting comprising automatically configuring ports of a TOR switch (e.g., TOR switches **105** and **157**) for each rack in the group of racks to provide access, by a respective server in each rack, to each network (e.g., virtual networks **121**) in the group of networks.

In one embodiment, a method comprises: receiving, over a network, a request to configure a first rack of at least one data center; receiving, over the network, configuration data

15

for at least one first virtual network to be accessed by a first computing device mounted in the first rack; and in response to receiving the configuration data, configuring a first TOR switch of the first rack, the configuring including associating at least one first virtual network with the first TOR switch.

In one embodiment, the method further comprises: receiving, from a user interface of a client device, a request to create a group of networks, the group including the at least one virtual network; in response to receiving the request, creating the group of networks; and in response to receiving a configuration selection made in the user interface, connecting the group of networks to a group of racks, the connecting comprising automatically configuring ports of a TOR switch for each rack in the group of racks to provide access, by a respective server in each rack, to each network in the group of networks.

In one embodiment, a system comprises: a network fabric to transmit data in at least one data center, wherein the at least one data center includes racks for mounting servers connected to the network fabric; network switches (e.g., TOR switches **105**, **157**) connected to the network fabric, wherein each network switch corresponds to a respective one of the racks; a data store (e.g., data store **131**) to store configuration data for the network switches; at least one processing device; and memory containing instructions configured to instruct the at least one processing device to: receive, via a portal from a client device, a request to create a computing environment supported on a plurality of racks (e.g., racks **103**, **155**) connected by the network fabric, the plurality of racks including a first rack for mounting a physical server configured to communicate with a physical server of a second rack in the computing environment; create at least one virtual network (e.g., virtual networks **121**) in the computing environment; and configure at least one of the network switches to associate the at least one virtual network with the physical server.

In one embodiment, the instructions are further configured to instruct the at least one processing device to: receive, via a user interface of the client device, configuration selections associated with a new computing device in the computing environment; based on the configuration selections, configure processing resources, memory resources, and storage resources; and deploy the new computing device to the computing environment, wherein the new computing device is configured to run a virtual server connected to the at least one virtual network.

FIG. 6 shows a block diagram of a computing device, which can be used in various embodiments. While FIG. 6 illustrates various components, it is not intended to represent any particular architecture or manner of interconnecting the components. Other systems that have fewer or more components may also be used. In one embodiment, the computing device is a server. In one embodiment, several servers may be used and each reside on separate computing systems, or one or more may run on the same computing device, in various combinations.

In FIG. 6, computing device **8201** includes an inter-connect **8202** (e.g., bus and system core logic), which interconnects a microprocessor(s) **8203** and memory **8208**. The microprocessor **8203** is coupled to cache memory **8204** in the example of FIG. 6.

The inter-connect **8202** interconnects the microprocessor(s) **8203** and the memory **8208** together and also interconnects them to a display controller and display device **8207** and to peripheral devices such as input/output (I/O) devices **8205** through an input/output controller(s) **8206**. Typical I/O devices include mice, keyboards, modems,

16

network interfaces, printers, scanners, video cameras and other devices which are well known in the art.

The inter-connect **8202** may include one or more buses connected to one another through various bridges, controllers and/or adapters. In one embodiment the I/O controller **8206** includes a USB (Universal Serial Bus) adapter for controlling USB peripherals, and/or an IEEE-1394 bus adapter for controlling IEEE-1394 peripherals.

The memory **8208** may include ROM (Read Only Memory), and volatile RAM (Random Access Memory) and non-volatile memory, such as hard drive, flash memory, etc.

Volatile RAM is typically implemented as dynamic RAM (DRAM) which requires power continually in order to refresh or maintain the data in the memory. Nonvolatile memory is typically a solid-state drive, magnetic hard drive, a magnetic optical drive, or an optical drive (e.g., a DVD RAM), or other type of memory system which maintains data even after power is removed from the system. The non-volatile memory may also be a random access memory.

The non-volatile memory can be a local device coupled directly to the rest of the components in the computing device. A non-volatile memory that is remote from the computing device, such as a network storage device coupled to the computing device through a network interface such as a modem or Ethernet interface, can also be used.

In one embodiment, a computing device as illustrated in FIG. 6 is used to implement computing device **115**, computing device **125**, TOR switch **105**, server **107**, and/or other servers.

In another embodiment, a computing device as illustrated in FIG. 6 is used to implement a user terminal or a mobile device on which an application is installed or being installed. A user terminal may be in the form of, for example, a laptop or notebook computer, or a personal desktop computer.

In some embodiments, one or more servers can be replaced with the service of a peer to peer network of a plurality of data processing systems, or a network of distributed computing systems. The peer to peer network, or a distributed computing system, can be collectively viewed as a computing device.

Embodiments of the disclosure can be implemented via the microprocessor(s) **8203** and/or the memory **8208**. For example, the functionalities described can be partially implemented via hardware logic in the microprocessor(s) **8203** and partially using the instructions stored in the memory **8208**. Some embodiments are implemented using the microprocessor(s) **8203** without additional instructions stored in the memory **8208**. Some embodiments are implemented using the instructions stored in the memory **8208** for execution by one or more general purpose microprocessor(s) **8203**. Thus, the disclosure is not limited to a specific configuration of hardware and/or software.

FIG. 7 shows a block diagram of a computing device, according to one embodiment. In one embodiment, the computing device of FIG. 7 is used to implement client device **137**. In FIG. 7, the computing device includes an inter-connect **9221** connecting the presentation device **9229**, user input device **9231**, a processor **9233**, a memory **9227**, a position identification unit **9225** and a communication device **9223**.

In FIG. 7, the position identification unit **9225** is used to identify a geographic location. The position identification unit **9225** may include a satellite positioning system receiver, such as a Global Positioning System (GPS) receiver, to automatically identify the current position of the computing device.

In FIG. 7, the communication device 9223 is configured to communicate with a server to provide data, including configuration data and/or an image from a camera of the computing device. In one embodiment, the user input device 9231 is configured to receive or generate user data or content. The user input device 9231 may include a text input device, a still image camera, a video camera, and/or a sound recorder, etc.

Closing

The disclosure includes various devices which perform the methods and implement the systems described above, including data processing systems which perform these methods, and computer readable media containing instructions which when executed on data processing systems cause the systems to perform these methods.

The description and drawings are illustrative and are not to be construed as limiting. Numerous specific details are described to provide a thorough understanding. However, in certain instances, well-known or conventional details are not described in order to avoid obscuring the description. References to one or an embodiment in the present disclosure are not necessarily references to the same embodiment; and, such references mean at least one.

Reference in this specification to “one embodiment” or “an embodiment” means that a particular feature, structure, or characteristic described in connection with the embodiment is included in at least one embodiment of the disclosure. The appearances of the phrase “in one embodiment” in various places in the specification are not necessarily all referring to the same embodiment, nor are separate or alternative embodiments mutually exclusive of other embodiments. Moreover, various features are described which may be exhibited by some embodiments and not by others. Similarly, various requirements are described which may be requirements for some embodiments but not other embodiments.

As used herein, “coupled to” generally refers to a connection between components, which can be an indirect communicative connection or direct communicative connection (e.g., without intervening components), whether wired or wireless, including connections such as electrical, optical, magnetic, etc.

In this description, various functions and operations may be described as being performed by or caused by software code to simplify description. However, those skilled in the art will recognize what is meant by such expressions is that the functions result from execution of the code by one or more processors, such as a microprocessor, Application-Specific Integrated Circuit (ASIC), graphics processor, and/or a Field-Programmable Gate Array (FPGA). Alternatively, or in combination, the functions and operations can be implemented using special purpose circuitry (e.g., logic circuitry), with or without software instructions. Embodiments can be implemented using hardwired circuitry without software instructions, or in combination with software instructions. Thus, the techniques are not limited to any specific combination of hardware circuitry and software, nor to any particular source for the instructions executed by a computing device.

While some embodiments can be implemented in fully functioning computers and computer systems, various embodiments are capable of being distributed as a computing product in a variety of forms and are capable of being

applied regardless of the particular type of machine or computer-readable media used to actually effect the distribution.

At least some aspects disclosed can be embodied, at least in part, in software. That is, the techniques may be carried out in a computing device or other system in response to its processor, such as a microprocessor, executing sequences of instructions contained in a memory, such as ROM, volatile RAM, non-volatile memory, cache or a remote storage device.

Routines executed to implement the embodiments may be implemented as part of an operating system, middleware, service delivery platform, SDK (Software Development Kit) component, web services, or other specific application, component, program, object, module or sequence of instructions referred to as “computer programs.” Invocation interfaces to these routines can be exposed to a software development community as an API (Application Programming Interface). The computer programs typically comprise one or more instructions set at various times in various memory and storage devices in a computer, and that, when read and executed by one or more processors in a computer, cause the computer to perform operations necessary to execute elements involving the various aspects.

A machine readable medium can be used to store software and data which when executed by a computing device causes the device to perform various methods. The executable software and data may be stored in various places including, for example, ROM, volatile RAM, non-volatile memory and/or cache. Portions of this software and/or data may be stored in any one of these storage devices. Further, the data and instructions can be obtained from centralized servers or peer to peer networks. Different portions of the data and instructions can be obtained from different centralized servers and/or peer to peer networks at different times and in different communication sessions or in a same communication session. The data and instructions can be obtained in entirety prior to the execution of the applications. Alternatively, portions of the data and instructions can be obtained dynamically, just in time, when needed for execution. Thus, it is not required that the data and instructions be on a machine readable medium in entirety at a particular instance of time.

Examples of computer-readable media include but are not limited to recordable and non-recordable type media such as volatile and non-volatile memory devices, read only memory (ROM), random access memory (RAM), flash memory devices, solid-state drive storage media, removable disks, magnetic disk storage media, optical storage media (e.g., Compact Disk Read-Only Memory (CD ROMs), Digital Versatile Disks (DVDs), etc.), among others. The computer-readable media may store the instructions.

In general, a tangible or non-transitory machine readable medium includes any mechanism that provides (e.g., stores) information in a form accessible by a machine (e.g., a computer, mobile device, network device, personal digital assistant, manufacturing tool, any device with a set of one or more processors, etc.).

In various embodiments, hardwired circuitry may be used in combination with software instructions to implement the techniques. Thus, the techniques are neither limited to any specific combination of hardware circuitry and software nor to any particular source for the instructions executed by a computing device.

Although some of the drawings illustrate a number of operations in a particular order, operations which are not order dependent may be reordered and other operations may

be combined or broken out. While some reordering or other groupings are specifically mentioned, others will be apparent to those of ordinary skill in the art and so do not present an exhaustive list of alternatives. Moreover, it should be recognized that the stages could be implemented in hardware, firmware, software or any combination thereof.

In the foregoing specification, the disclosure has been described with reference to specific exemplary embodiments thereof. It will be evident that various modifications may be made thereto without departing from the broader spirit and scope as set forth in the following claims. The specification and drawings are, accordingly, to be regarded in an illustrative sense rather than a restrictive sense.

Various embodiments set forth herein can be implemented using a wide variety of different types of computing devices. As used herein, examples of a “computing device” include, but are not limited to, a server, a centralized computing platform, a system of multiple computing processors and/or components, a mobile device, a user terminal, a vehicle, a personal communications device, a wearable digital device, an electronic kiosk, a general purpose computer, an electronic document reader, a tablet, a laptop computer, a smartphone, a digital camera, a residential domestic appliance, a television, or a digital music player. Additional examples of computing devices include devices that are part of what is called “the internet of things” (IOT). Such “things” may have occasional interactions with their owners or administrators, who may monitor the things or modify settings on these things. In some cases, such owners or administrators play the role of users with respect to the “thing” devices. In some examples, the primary mobile device (e.g., an Apple iPhone) of a user may be an administrator server with respect to a paired “thing” device that is worn by the user (e.g., an Apple watch).

In some embodiments, the computing device can be a host system, which is implemented, for example, as a desktop computer, laptop computer, network server, mobile device, or other computing device that includes a memory and a processing device. The host system can include or be coupled to a memory sub-system so that the host system can read data from or write data to the memory sub-system. The host system can be coupled to the memory sub-system via a physical host interface.

Examples of a physical host interface include, but are not limited to, a serial advanced technology attachment (SATA) interface, a peripheral component interconnect express (PCIe) interface, universal serial bus (USB) interface, Fibre Channel, Serial Attached SCSI (SAS), a double data rate (DDR) memory bus, etc. The physical host interface can be used to transmit data between the host system and the memory sub-system. The host system can further utilize an NVM Express (NVMe) interface to access memory components of the memory sub-system when the memory sub-system is coupled with the host system by the PCIe interface. The physical host interface can provide an interface for passing control, address, data, and other signals between the memory sub-system and the host system. In general, the host system can access multiple memory sub-systems via a same communication connection, multiple separate communication connections, and/or a combination of communication connections.

In one embodiment, the host system includes a processing device and a controller. The processing device of the host system can be, for example, a microprocessor, a graphics processing unit, a central processing unit (CPU), an FPGA, a processing core of a processor, an execution unit, etc. In one example, the processing device can be a single package

that combines an FPGA and a microprocessor, in which the microprocessor does most of the processing, but passes off certain predetermined, specific tasks to an FPGA block. In one example, the processing device is a soft microprocessor (also sometimes called softcore microprocessor or a soft processor), which is a microprocessor core implemented using logic synthesis. The soft microprocessor can be implemented via different semiconductor devices containing programmable logic (e.g., ASIC, FPGA, or CPLD).

In some examples, the controller is a memory controller, a memory management unit, and/or an initiator. In one example, the controller controls the communications over a bus coupled between the host system and the memory subsystem.

In general, the controller can send commands or requests to the memory subsystem for desired access to the memory components. The controller can further include interface circuitry to communicate with the memory sub-system. The interface circuitry can convert responses received from the memory sub-system into information for the host system. The controller of the host system can communicate with the controller of the memory sub-system to perform operations such as reading data, writing data, or erasing data at the memory components and other such operations.

In some instances, a controller can be integrated within the same package as the processing device. In other instances, the controller is separate from the package of the processing device. The controller and/or the processing device can include hardware such as one or more integrated circuits and/or discrete components, a buffer memory, a cache memory, or a combination thereof. The controller and/or the processing device can be a microcontroller, special purpose logic circuitry (e.g., a field programmable gate array (FPGA), an application specific integrated circuit (ASIC), etc.), or another suitable processor.

The memory components can include any combination of the different types of non-volatile memory components and/or volatile memory components. An example of non-volatile memory components includes a negative-and (NAND) type flash memory. Each of the memory components can include one or more arrays of memory cells such as single level cells (SLCs) or multi-level cells (MLCs) (e.g., triple level cells (TLCs) or quad-level cells (QLCs)). In some embodiments, a particular memory component can include both an SLC portion and a MLC portion of memory cells. Each of the memory cells can store one or more bits of data (e.g., data blocks) used by the host system. Although non-volatile memory components such as NAND type flash memory are described, the memory components can be based on any other type of memory such as a volatile memory.

In some embodiments, the memory components can be, but are not limited to, random access memory (RAM), read-only memory (ROM), dynamic random access memory (DRAM), synchronous dynamic random access memory (SDRAM), phase change memory (PCM), magnetoresistive random access memory (MRAM), Spin Transfer Torque (STT)-MRAM, ferroelectric random-access memory (FeTRAM), ferroelectric RAM (FeRAM), conductive bridging RAM (CBRAM), resistive random access memory (RRAM), oxide based RRAM (OxRAM), negative-or (NOR) flash memory, electrically erasable programmable read-only memory (EEPROM), nanowire-based non-volatile memory, memory that incorporates memristor technology, and a cross-point array of non-volatile memory cells. A cross-point array of non-volatile memory can perform bit storage based on a change of bulk resistance, in conjunction with a

21

stackable crossgridded data access array. Additionally, in contrast to many flash-based memories, cross-point non-volatile memory can perform a write in-place operation, where a nonvolatile memory cell can be programmed without the non-volatile memory cell being previously erased. Furthermore, the memory cells of the memory components can be grouped as memory pages or data blocks that can refer to a unit of the memory component used to store data.

The controller of the memory sub-system can communicate with the memory components to perform operations such as reading data, writing data, or erasing data at the memory components and other such operations (e.g., in response to commands scheduled on a command bus by a controller). A controller can include a processing device (processor) configured to execute instructions stored in local memory. The local memory of the controller can include an embedded memory configured to store instructions for performing various processes, operations, logic flows, and routines that control operation of the memory sub-system, including handling communications between the memory sub-system and the host system. In some embodiments, the local memory can include memory registers storing memory pointers, fetched data, etc. The local memory can also include read-only memory (ROM) for storing micro-code. While the example memory sub-system includes the controller, in another embodiment of the present disclosure, a memory sub-system may not include a controller, and can instead rely upon external control (e.g., provided by an external host, or by a processor or controller separate from the memory sub-system).

In general, the controller can receive commands or operations from the host system and can convert the commands or operations into instructions or appropriate commands to achieve the desired access to the memory components. The controller can be responsible for other operations such as wear leveling operations, garbage collection operations, error detection and error-correcting code (ECC) operations, encryption operations, caching operations, and address translations between a logical block address and a physical block address that are associated with the memory components. The controller can further include host interface circuitry to communicate with the host system via the physical host interface. The host interface circuitry can convert the commands received from the host system into command instructions to access the memory components as well as convert responses associated with the memory components into information for the host system.

The memory sub-system can also include additional circuitry or components that are not illustrated. In some embodiments, the memory sub-system can include a cache or buffer (e.g., DRAM or SRAM) and address circuitry (e.g., a row decoder and a column decoder) that can receive an address from the controller and decode the address to access the memory components.

What is claimed is:

1. A computing system comprising:

a first server rack and a second server rack located in different locations, the first server rack containing a first top-of-rack (TOR) switch connected to a first server, the second server rack containing a second TOR switch connected to a second server;

a first computing device separate from the first server rack and the second server rack and configured to:

receive, from a client device accessed by a first user and configured to be in communication with the first computing device, a request to create a group of a plurality of virtual networks, to implement a virtual

22

extensible local area network (VXLAN), and to connect a virtual server to the group;
create the plurality of virtual networks in response to the request;

assign the plurality of virtual networks to the group;
configure at least one switch of a network fabric to associate the virtual server with the plurality of virtual networks of the group;

in response to receiving a configuration selection made via a user interface, connect the group of the plurality of virtual networks to the first and second server racks by automatically configuring ports of the first and second TOR switches to provide access, by at least one server of the first and second server racks, to each virtual network of the plurality of virtual networks of the group;

implement the VXLAN for transmitting data between the first and second server based on the request from the client device accessed by the first user and configured to be in communication with the first computing device by converting at least one virtual networks of the group of the plurality of virtual networks into the VXLAN, wherein the first computing device includes an Internet configuration manager for configuring IP services to provide Internet connectivity to the first and second server based on the request from the client device, and wherein the Internet configuration manager is configured to communicate configuration data for configuring the first and second TOR switch to connect the first and second server to the IP services;

a second computing device separate from the first server rack and the second server rack and configured to allocate the second server to handle workloads of the first user by forming a virtual machine containing the first and second server at the request from the client device accessed by the first user;

wherein the second computing device executes a hyper-converged management service to allocate compute, memory and storage resources provided by the first server rack and the second server rack and stores data regarding the allocation of resources in a second data store; and

the network fabric,

wherein the first computing device and the second computing device are coupled in electronic communication with the first server rack and the second server rack via the network fabric.

2. The computing system of claim 1, further comprising a first and second physical fiber port coupling the first and second server rack, respectively, to the network fabric.

3. The computing system of claim 1, wherein the first TOR switch includes a memory and a plurality of ports, wherein the memory is configured to store connection assignments of the plurality of ports to a router and servers in the first server rack.

4. The computing system of claim 1, wherein the first computing device includes a switch configuration manager to implement the VXLAN and a first data store to store configurations of the switch configuration manager.

5. The computing system of claim 1, wherein the client device communicates with the first computing device by an API through an external network.

6. The computing system of claim 1, wherein the network fabric includes the plurality of virtual networks to be selected by the client device to be associated with the first and second server.

23

7. A computing system comprising:
- a first server rack and a second server rack located in different locations, the first server rack containing a first top-of-rack (TOR) switch connected to a first server, the second server rack containing a second TOR switch connected to a second server;
 - a first computing device separate from the first server rack and the second server rack and configured to:
 - receive, from a client device accessed by a first user and configured to be in communication with the first computing device, a request to create a group of a plurality of virtual networks, to implement a virtual extensible local area network (VXLAN), and to connect a virtual server to the group;
 - create the plurality of virtual networks in response to the request;
 - assign the plurality of virtual networks to the group;
 - configure at least one switch of a network fabric to associate the virtual server with the plurality of virtual networks of the group;
 - in response to receiving a configuration selection made via a user interface, connecting the group of the plurality of virtual networks to the first and second server racks by automatically configuring ports of the first and second TOR switches to provide access, by at least one server of the first and second server racks, to each virtual network of the plurality of virtual networks of the group;
 - a switch configuration manager configured to implement the virtual extensible local area network (VXLAN) for transmitting data between the first and second server at the request from the client device accessed by a first user by converting at least one virtual networks of the group of the plurality of virtual networks into the VXLAN;
 - a hyper-converged management service configured to allocate the second server to handle workloads of the first user by forming a virtual machine containing the first and second server at the request from the client device accessed by the first user, wherein the hyper-converged management service is installed on the first computing device separate from the first server rack and the second server rack;
 - the network fabric coupled to the first and second server rack, the switch configuration manager and the hyper-converged management service;
 - a first and second physical fiber port coupling the first and second server rack, respectively, to the network fabric; and
 - an Internet configuration manager for configuring IP services to provide Internet connectivity to the first and second server at a request from the client device, wherein the Internet configuration manager is configured to communicate configuration data needed by the switch configuration manager for configuring the first and second TOR switch to connect the first and second server to the IP services.
8. The computing system of claim 7, further comprising a first and second physical fiber port coupling the first and second server rack, respectively, to the network fabric.
9. The computing system of claim 7, wherein the first TOR switch includes a memory and a plurality of ports, wherein the memory is configured to store connection assignments of the plurality of ports to a router and servers in the first server rack.

24

10. The computing system of claim 7, wherein the switch configuration manager is configured to store configurations of the switch configuration manager in a first data store.
11. The computing system of claim 7, wherein the client device communicates with the switch configuration manager by an API through an external network.
12. The computing system of claim 7, wherein the hyper-converged management service allocates compute, memory and storage resources provided by the first and second server rack and stores data regarding the allocation of resources in a second data store.
13. The computing system of claim 7, wherein the network fabric includes the plurality of virtual networks to be selected by the client device to be associated with the first and second server.
14. A method for providing access to multiple server racks, the method comprising:
- receiving a request from a client device accessed by a first user to allocate more computing resources, to create a group of a plurality of virtual networks, to implement a virtual extensible local area network (VXLAN), and to connect a virtual server to the group;
 - creating the plurality of virtual networks in response to the request;
 - assigning the plurality of virtual networks to the group;
 - configuring at least one switch of a network fabric to associate the virtual server with the plurality of virtual networks of the group;
 - in response to receiving a configuration selection made via a user interface, connecting the group of the plurality of virtual networks to first and second server racks by automatically configuring ports of a first top-of-rack (TOR) switch and a second TOR switch to provide access, by at least one server of the first and second server racks, to each virtual network of the plurality of virtual networks of the group;
 - connecting, upon receiving the request, a first server contained in the first server rack with a second server contained in the second server rack through the first TOR switch contained in the first server rack and the second TOR switch contained in the second server rack, the first and second server rack being located in different locations;
 - implementing, by a first computing device, the VXLAN for transmitting data between the first and second server at the request from the client device accessed by the first user by converting at least one virtual networks of the group of the plurality of virtual networks into the VXLAN;
 - configuring, by the first computing device, IP services to provide Internet connectivity to the first and second server at a request from the client device, wherein the IP services correspond to a public IP address and protocol selected using the user interface of the client device;
 - allocating, by a second computing device separate from the first server rack and the second server rack, the second server to handle workloads of the first user by forming a virtual machine containing the first and second server at the request from the client device; wherein the second computing device executes a hyper-converged management service to allocate compute, memory and storage resources provided by the first and second server rack and stores data regarding the allocation of resources in a second data store; and

25

providing the network fabric coupled to the first and second server rack and the first and second computing device,

wherein the first computing device and the second computing device are coupled in electronic communication with the first and second server rack via the network fabric. 5

15. The method of claim **14**, further comprising providing a first and second physical fiber port coupling the first and second server rack, respectively, to the network fabric. 10

16. The method of claim **14**, wherein the network fabric includes the plurality of virtual networks to be selected by the client device to be associated with the first and second server.

* * * * *

15

26